

Cost implications of cluster plot design choices for precise estimation of forest attributes in landscapes and forests of varying heterogeneity

Andrew J. Lister and Laura P. Leites

Abstract: Tradeoffs occur when deciding between improving forest inventory precision by increasing sample size or by augmenting cluster plot design factors like size or subplot separation distance. The nature of these tradeoffs changes with variation in type and scale of the spatial pattern of the attribute of interest. To understand the impacts of relationships between type and scale of spatial heterogeneity and cluster plot design efficiency, we constructed a factorial simulation experiment and analysed relationships among forest inventory cost, cluster plot design factors, and different spatial heterogeneity scenarios constructed via simulation. To calculate cost, we constructed a cost model that accounted for both on- and between-plot costs. We found that type and scale of heterogeneity have important implications for plot design choices. Homogeneous stands and landscapes are the least costly to inventory. Subplot area and count have stronger impacts than subplot separation on cost efficiency, particularly in landscapes with aggregated forest patterns and in stands with homogeneous tree patterns. We discuss results in the context of the physical interaction between cluster plot geometry and spatial patterns at different scales, provide computer code for simulations, and suggest principles that forest inventory cluster plot design specialists should consider when designing inventories.

Key words: cluster plot design, forest inventory design optimization, forest inventory cost model, forest pattern simulation, forest sampling simulation.

Résumé : Pour améliorer la précision d'un inventaire forestier, des compromis sont possibles entre l'augmentation de la taille de l'échantillon ou l'augmentation des facteurs de déploiement des placettes en grappes comme la taille des sous-placettes ou la distance qui les séparent. La nature de ces compromis change selon la variation du type et de l'échelle de l'arrangement spatial de l'attribut d'intérêt. Afin de comprendre l'impact des relations entre l'efficacité du déploiement des placettes en grappes et le type et l'échelle de l'hétérogénéité spatiale, nous avons réalisé une expérience de simulation factorielle et avons analysé les relations entre le coût de l'inventaire forestier, les facteurs de déploiement des placettes en grappes et différents scénarios d'hétérogénéité spatiale conçus par simulation. Pour calculer les coûts, nous avons conçu un modèle de coûts qui tient compte à la fois des coûts à l'intérieur et à l'extérieur des placettes. Nous avons constaté que le type et l'échelle de l'hétérogénéité ont des implications importantes pour le choix du déploiement des placettes. Les peuplements et les paysages homogènes sont les moins coûteux à inventorier. La superficie et le nombre de sous-placettes ont de plus grands impacts sur la rentabilité que la distance entre les sous-placettes, en particulier dans les paysages avec des arrangements spatiaux de peuplements regroupés et dans les peuplements avec des arrangements spatiaux d'arbres homogènes. Nous discutons des résultats dans le contexte de l'interaction physique entre la géométrie des placettes en grappes et les arrangements spatiaux à différentes échelles. De plus, nous fournissons un programme informatique pour les simulations et suggérons des principes dont les spécialistes du déploiement des placettes d'inventaire forestier en grappes devraient tenir compte lors de la conception des inventaires. [Traduit par la Rédaction]

Mots-clés : déploiement des placettes en grappes, optimisation de la conception de l'inventaire forestier, modèle de coûts d'inventaire forestier, simulation de l'arrangement spatial des forêts, simulation de l'échantillonnage forestier.

Introduction

Forest inventories are becoming increasingly important as countries seek to generate information for policy and management decisions, as well as participate in incentives programs aimed at reducing carbon emissions (IPCC 2006; GFOI 2016; FAO 2017). Technical and logistical challenges occur when designing inventories that cover a mosaic of areas with different types and scales of spatial heterogeneity. Clear, science-based guiding principles are therefore needed to develop inventory systems that cost-effectively meet users' needs and function across a broad range of tree and forest spatial pattern conditions.

There are several existing frameworks that have been proposed for designing efficient forest inventories. The most common design approach is constructing the inventory to meet precision constraints for the least cost (Freese 1962; Wiant and Yandle 1980; Lynch 2017a). Precision constraints, commonly termed allowable error (AE), are typically presented as the desired confidence interval, sampling error, or margin of error of the estimate that will arise from the inventory, along with an associated confidence level. For example, the Intergovernmental Panel on Climate Change (IPCC) recommends an AE of 10% at the 95% confidence level for defensible estimates of carbon stock change (IPCC 2006). AE is

Received 1 December 2020. Accepted 3 June 2021.

A.J. Lister. USDA Forest Service, Northern Research Station, Forest Inventory and Analysis, 3460 Industrial Drive, York, PA 17402, USA.

L.P. Leites. Department of Ecosystem Science and Management, Penn State University, 312 Forest Resources Building, University Park, PA 16802, USA.

Corresponding author: Andrew J. Lister (email: andrew.lister@usda.gov).

© 2021 The Author(s). Permission for reuse (free in most cases) can be obtained from copyright.com.

used in a straightforward manner in the inventory design process, through the following equation:

$$(1) \quad n_{\text{req}} = \left(\frac{CV\% \times t}{AE\%} \right)^2$$

where n_{req} is the sample size required to achieve a specified percent standard error (AE%), CV is a known or hypothesized coefficient of variation (CV) of the attribute of interest, and t is the Student's t value associated with the chosen confidence level (Loetsch and Haller 1973).

A key challenge in choosing optimal inventory and plot designs is that relationships between precision of the estimate and plot design change under different attribute spatial heterogeneity scenarios, and it is difficult to predict how these changes will occur (Kleinn 1994; Yim et al. 2015). Sources of heterogeneity include fine-scale tree patterns (Stoyan and Penttinen 2000), as well as forest patch patterns at the landscape scale (Riitters et al. 2016). Many forest inventories, like the national forest inventory of the United States, distribute plots across all lands in a spatially balanced way, including in nonforest areas, since one goal is to monitor forest area dynamics (Bechtold and Patterson 2005). Nonforest values under this paradigm are therefore included as valid plot-level measurements of zero for the attribute of interest (e.g., forest tree density or volume); this contributes, along with variance from measurements within forest patches, to the overall variance of the estimates.

The use of cluster plots further complicates forest inventory design decisions. In single stage cluster designs, primary units (sampling units, often referred to as plots) are composed of more than one secondary unit (measurement units, often referred to as subplots) distributed in a consistent pattern such as a line, cross, or L-shape. Cluster plots are generally more efficient because, for the same area measured, they can potentially capture more of the variability that is found on the landscape by avoiding sampling redundant information in spatially autocorrelated patches (Thompson 2012). Cluster plot design parameters include subplot size, count, and separation distance. Varying these parameters affects inventory costs, such as those accrued during plot and subplot establishment and walking between subplots.

The effect of each of these design parameters on inventory efficiency is hard to predict because of the interaction of different design variables and environmental heterogeneity scale and type. Impacts of cluster plot design on precision have previously been studied using data from mapped stands (Schreuder et al. 1987; Zenner and Peck 2009; Picard et al. 2018), existing forest inventory plots (e.g., Lynch 2003; Picard et al. 2004; Yim et al. 2015), and simulated forest tree patterns (Arvanitis and O'Regan 1967; Mackisack and Wood 1990; Brink and Schreuder 1992; Hou et al. 2015; Gove 2017). Based on our review of the literature, however, there are no studies that use a factorial simulation experiment to model the relationship among plot design parameters, spatial heterogeneity type and scale, and inventory efficiency and costs. Such an experimental approach would allow for control over plot design and heterogeneity scenario parameters, without confounding factors that are common in studies that rely wholly on field measurements. This would increase understanding of the complex relationships among these variables and provide insight into the mechanisms that govern their impacts on inventory efficiency.

In this study, we modelled inventory cost as a function of landscape- and stand-scale heterogeneity factors and cluster plot design variables including subplot area, count of subplots, and subplot separation distance. The main goal was to reveal how combinations of these factors and their interactions with scale and type of heterogeneity affect cost, and to provide insights that can help guide the development of efficient inventory networks under different heterogeneity scenarios.

Methods

Simulation experiment

A repeated measures factorial simulation experiment with multiple crossed factors (Oehlert 2000, p. 438) was conducted to model forest inventory cost as a function of two spatial heterogeneity factors and three cluster plot design factors. Each of the two spatial pattern factors, representing landscape (L) and stand (S) scale heterogeneity, had three levels (Fig. 1):

L1: highly dispersed patterns of forest patches with many small, isolated fragments

L2: intermediate levels of aggregation

L3: highly aggregated patterns, with large, continuous forest patches

S1: highly dispersed (uniformly distributed) pattern of tree locations

S2: completely random pattern of tree locations

S3: highly aggregated pattern of tree locations, with trees occurring in clumps

For each of the nine factor level combinations, 30 replicates were generated via simulation for a total of 270 replicates (procedure described in the section on simulation details below and in Fig. 1).

Each of the 270 replicate heterogeneity scenarios was sampled once using systematic sampling with 49 plot locations and a cluster plot design arising from every combination of the following factors:

d: 10, 25, or 50 m subplot separation (measured between subplot edges)

m: 2, 3, 4, or 5 subplots per plot

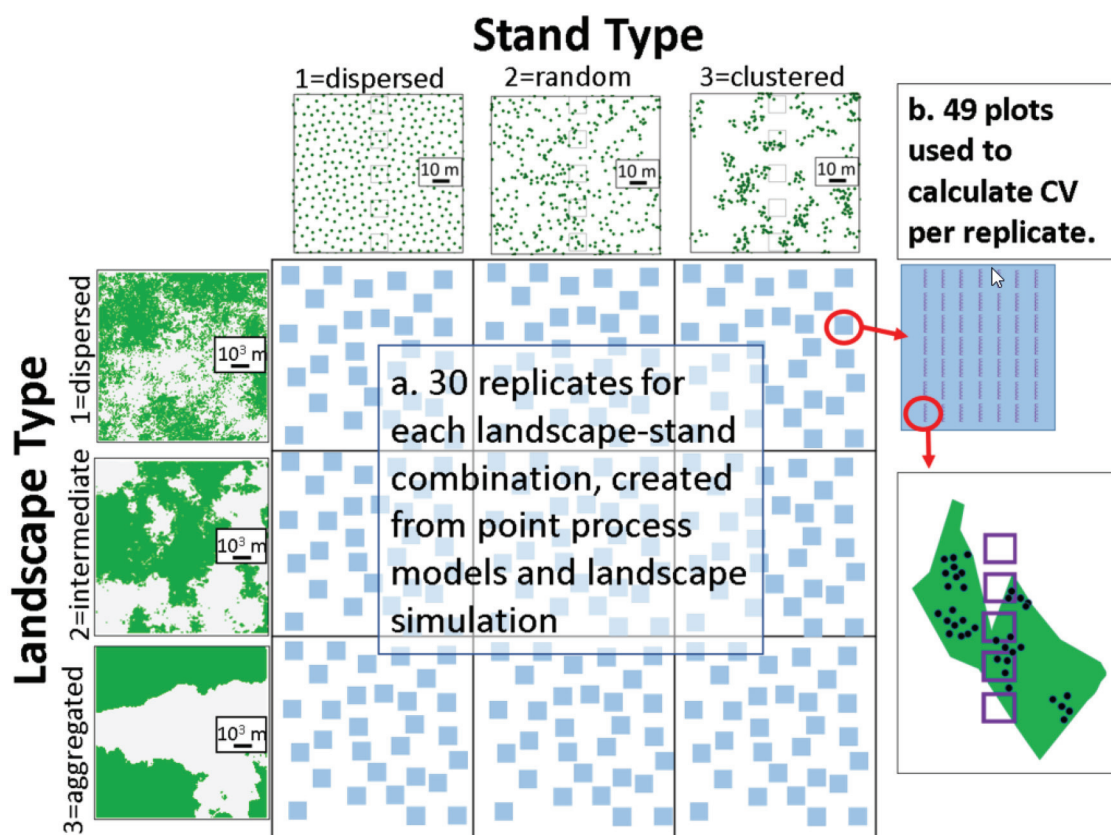
a: 0.01 or 0.2 ha subplot area

In other words, each member of a set of $3(d) \times 4(m) \times 2(a) = 24$ cluster plot designs was used to sample each of the 270 replicates at the 49 plot locations, as shown conceptually in Fig. 1. This process required a repeated measurements design, as the 49 plot locations were always the same on each replicate; the only thing that changed was the cluster plot design.

The variable recorded for each of the $30 \times 3(L) \times 3(S) \times 3(d) \times 4(m) \times 2(a) = 6480$ observations was CV of trees per hectare (herein-after N), calculated from the 49 plot-level values and using simple random sampling estimators under a single stage cluster sampling paradigm (Thompson 2012). CV was chosen because it is needed to estimate the required sample sizes (n_{req} ; eq. 1). Required sample size and plot design parameter costs were used to calculate total inventory cost, the dependent variable we used in our models, for each of the 6480 samples. A summary of this experiment is as follows:

1. Simulate 270 replicates of the forested heterogeneity scenarios (9 L - S combinations $\times$ 30 replicates each; Fig. 1a).
2. Overlay a 7×7 grid of plot locations on each replicate.
3. Sample each of the 270 replicates at the plot locations from step 2 using each of the 24 cluster plot designs, recording CV of N for each case (Fig. 1b), leading to 6480 observations of CV.
4. Use eq. 1 to calculate required sample size (n_{req}) for each observation to achieve an AE of 10% sampling error at the 95% level.
5. Using knowledge of n_{req} , landscape size, and various estimates of on- and between-plot costs, implement a cost model (described in the Cost Model section below and in Appendix A) to calculate total inventory cost for each observation.
6. Build and interpret a linear model of total inventory cost as a function of L , S , d , m , and a .

Fig. 1. Conceptual model depicting factorial experimental design. (a) Thirty replicates for each of three levels of landscape-scale (*L*) heterogeneity were simulated for each of three levels of stand-scale (*S*) heterogeneity. (b) At each plot located on a 7×7 grid superimposed on each landscape, stand-scale patterns were simulated for each level of *S*. Candidate cluster plot designs were superimposed over each plot location for each replicate, and the number of trees per hectare was calculated per plot. Coefficients of variation from the 49 plots were calculated for each combination of candidate plot design and replicate. [Colour online.]



Simulation procedure

Simulation of *L*

To create a heterogeneity gradient from simulated landscapes, we used the multifractal map generation feature of the qRule landscape analysis software (Gardner 2017, 1999). We created square maps (10.24 km sides, 105 km²) with 50% coverage of each of two landcover classes: forest and nonforest. The software generates realistic maps using a fractal algorithm that produces randomized, spatially correlated patterns of land cover, with the option of controlling the level of aggregation. This allows for the creation of each level of *L*, described above and in Fig. 1. We calibrated this algorithm such that the three levels of aggregation corresponded with a gradient of approximate forest edge densities of 432.5 m/ha for *L1*, 55.0 m/ha for *L2*, and 11.2 m/ha for *L3*. Forest edge density is defined as the length of the interface between forest and nonforest pixels, divided by the area of the map. We chose the proportion forest (50%) and edge density parameters to reflect a range of realistic landscape patterns such as those found in northeastern United States temperate forests that intermingle with agricultural and urban areas (e.g., *L2* and *L3*), and those found in silvopastoral ecosystems like those in Central America with sparse tree clusters of different sizes (*L1*). We provide examples of these types of landscapes and provide code for calculating edge density and forest proportion from existing maps in Supplementary Material S1a¹.

For each level of *L*, 90 maps (replicates) were simulated, 30 per level of *S*. The raster (Hijmans 2019) and spatstat (Baddeley and Turner 2005) R packages were used to convert the qRule output files to raster objects.

Simulation of *S*

For the simulations of tree patterns, an *N* of 388 trees/ha was chosen. This is the average tree density of live trees greater than or equal to 12.7 cm diameter at breast height on forest land for Pennsylvania, according to the USDA Forest Service's Forest Inventory and Analysis (FIA) database (USDA 2020). We considered other commonly reported inventory attributes to use for our study, including basal area per hectare (hereinafter *G*) and volume per hectare (hereinafter *V*). *N* was chosen because, in Pennsylvania, the variance of its estimate for trees greater than 12.7 cm diameter at breast height tends to be slightly higher than that for *G* but slightly lower than that for *V* (USDA 2020). In addition, *N* and *G* are co-equal components of stocking calculations, making it arguably one of the fundamental forest inventory attributes (Zarnoch and Bechtold 2000). Finally, simulating patterns of *N* is possible using standard point process models that reflect naturally occurring patterns, while simulating *G* and *V* adds complexity by requiring hierarchical models (Lister and Leites 2018) or approaches like the use of copulas (Kershaw et al. 2010) that incorporate effects of tree characteristics in spatial pattern generation.

¹Supplementary data are available with the article at <https://doi.org/10.1139/cjfr-2020-0509>.

Stoyan and Penttinen (2000) describe how various ecological conditions can lead to regular, random, or dispersed patterns of trees. Regular patterns would occur in plantations with a uniform spacing, or in ecosystems where competition has a repulsive effect on the establishment of individual trees. Random patterns can occur at various seral stages in different ecosystems; Lister and Leites (2018) found that mid-sized trees exhibited random patterns in a secondary growth forest in Pennsylvania, USA. Clustering can occur due to gap dynamics and patterns of mortality and recruitment (Reich and Arvanitis 1992; Larson et al. 2015). To create a heterogeneity gradient from tree spatial patterns at the stand scale, spatial point process models were used to simulate the three types of *S* pattern types described above at each plot location using the spatstat R package. To create the highly dispersed pattern (*S1*), a simple sequential inhibition pattern generator (rssi) with a 4 m inhibition distance (as might occur in a plantation) and the requisite number of points to achieve the target *N* was implemented. For the intermediate level of aggregation (*S2*), a homogeneous Poisson process pattern generator (rpoispp) was used with the target *N* to create a completely random spatial pattern (as might occur at various times in forest succession, or with mid-sized trees in some forests (Lister and Leites 2018)). To create the aggregated spatial pattern (*S3*), a Thomas cluster process generator (rThomas) was used with a scale parameter of 3 m, a mean number of points per cluster of 10, and the requisite *N* for cluster centers. The scale parameter affects the compactness of the clusters; to compare our parameter choice to those that might be found in natural forests, we conducted the analysis described in Supplementary Material S1b¹.

For each level of *S*, 90 realizations (replicates) were generated, 30 per level of *L*. For each of the 49 plots in each of the *L-S* replicates, simulations were performed in a rectangular, 144.7 m × 523.6 m window surrounding each cluster plot center. This window size, which represents a 50 m buffer around the largest candidate cluster plot design's footprint, was chosen to minimize artefacts in the point pattern simulation process that would occur from restricting the simulation to the plot boundaries.

Cluster plot design creation

At each plot location, cluster plots with different *d-m-a* configurations at each of 49 locations were superimposed over the simulated *L-S* combinations, as shown in Fig. 1b, using the spatstat, raster (Hijmans 2019), and sp (Pebesma and Bivand 2005) packages. Cluster plot shapes consisted of a linear array of square subplots aligned north to south (Fig. 1). Candidate cluster plot designs were located one at a time, and *N* for that cluster plot recorded.

We have provided R code in Supplementary Material S2¹ for generating the cluster plot designs and *L-S* scenarios, and for sampling each *L-S-d-m-a* combination. This code can be used to conduct trials with different combinations of the parameters and calibrate cost models that are appropriate to specific situations.

Cost model

The details of the cost model used are described in Appendix A. The cost modelling approach used in this study follows that of Zeide (1980) and Wiant and Yandle (1980), which is based principally on the proportions of on- and between-plot times. To make the proposed framework more accessible to practitioners, salary and daily cost rates are included in the cost model. This model is not meant to fully represent reality, but rather to serve as a tool with which to assign a reasonable total inventory cost (the dependent variable used in the linear model) to each of the 6480 factor combinations described above, taking into account on-plot, between-plot, and logistical costs. The components of the cost model include

Table 1. Parameters, *F* statistics, and *p* values associated with the final model from eq. 2.

Parameter	<i>F</i> value	<i>p</i> value
(Intercept)	24 584.40	<0.0001
a	267 293.70	<0.0001
a:L	18 747.34	<0.0001
a:m	13 262.61	<0.0001
m	3 541.28	<0.0001
a:m:L	1 063.06	<0.0001
L	873.66	<0.0001
m:L	683.10	<0.0001
a:S	497.44	<0.0001
S	272.30	<0.0001
L:d	79.48	<0.0001
m:S	60.04	<0.0001
a:d	30.27	<0.0001
a:m:S	26.61	<0.0001
D	7.17	0.0008
a:m:d	6.20	<0.0001
m:d	2.43	0.0239

- On-plot costs.** These are calculated for each cluster plot design – replicate combination based on assumptions about crew salaries and walking and measurement times required to measure all the trees in a given simulated pattern.
- Between-plot costs.** These are calculated using the time required to visit all plots for a given cluster plot design – spatial heterogeneity scenario, based on the size of the study area, travel speeds, crew salaries, and n_{req} (eq. 1).

Total inventory cost is calculated as the combination of on- and between-plot costs, as well as daily logistical costs associated with the length of time spent executing the inventory.

Analysis

Version 3.1-2 of the lmerTest package of version 3.5.1 of the R statistical software (Bates et al. 2015; Kuznetsova et al. 2017; R Core Team 2018) was used to fit the linear model of the form

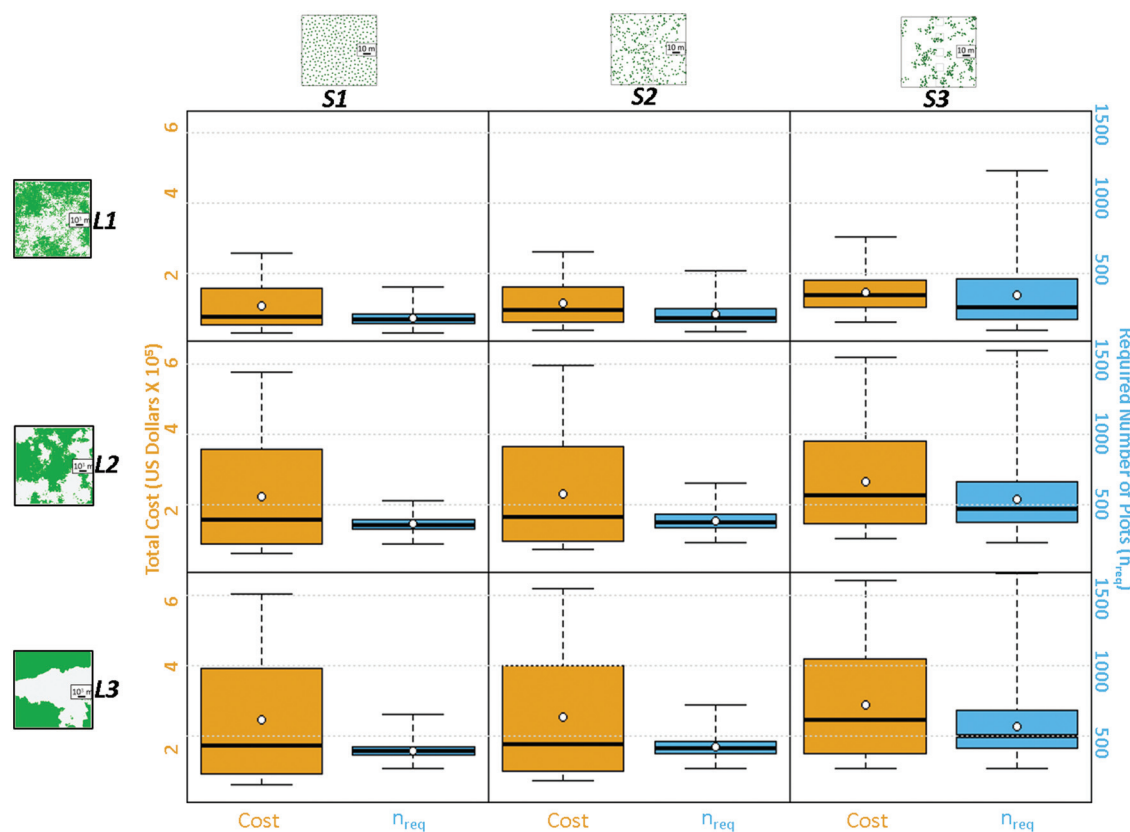
$$(2) \quad \text{Cost}_{ijkpq} = \mu + L_i \times S_j \times m_p \times d_q \times a_r + \gamma_{k(ij)} + \delta_{pk(ij)} + \delta_{qk(ij)} + \delta_{rk(ij)} + \varepsilon_{pqrk(ij)}$$

where Cost is the total inventory cost associated with the *k*th replicate, for the *i*th and *j*th levels of *L* and *S*, respectively, and for the *p*th, *q*th, and *r*th levels of the plot design variables *m*, *d*, and *a* respectively; μ is the overall mean; γ is the replicate random effect with $k = 1$ to 30 replicates; *L* is the landscape heterogeneity class with $i = 1$ to 3 levels; *S* is the stand heterogeneity class with $j = 1$ to 3 levels; *m* is the number of subplots with $p = 1$ to 4 levels; *d* is the distance between subplots with $q = 1$ to 3 levels; *a* is the subplot area with $r = 1$ to 2 levels; $\delta_{pk(ij)}$, $\delta_{qk(ij)}$, and $\delta_{rk(ij)}$ are random effects accounting for repeated measures for *m*, *d* and *a*, respectively; and $\varepsilon_{pqrk(ij)} \sim N(0, \sigma^2)$ is the whole model error term. After the full model was fit, model terms that were not significant were removed, the model was fit again, and modelling assumptions were evaluated using diagnostic graphics.

Results

Linear modelling results, which consider the repeated measures factorial design, indicate that landscape- and stand-level heterogeneity scale and type, as well as plot design factors, significantly affect cost, alone and in several interactions (Table 1). Model diagnostic graphics are provided in Supplementary Material S3¹. In the following sections, we present a summary of the effect of each main factor on cost by averaging across the other factors' levels,

Fig. 2. Boxplot of the total inventory cost (orange shading) and required number of plots (n_{req} , blue shading) values for each landscape and stand heterogeneity type combination. Shaded boxes are interquartile ranges, dark lines are medians, small circles are means, and dashed lines are range of values for each level. $n = 2160$ for each level. [Colour online.]



highlight the important interactions found, and then characterize the relationship between cost and n_{req} .

Impacts of L and S on costs and n_{req}

Landscape-scale heterogeneity L has a significant effect on inventory cost (Table 1). The more aggregated landscape patterns ($L2$ and $L3$) both require a higher (by 102% and 122%, respectively) average cost and higher (by 74% and 85%, respectively) average n_{req} to achieve AE than does the dispersed landscape type ($L1$; Fig. 2). The level of variability of the cost values (interquartile range (IQR) and range) suggests that plot design decisions have an important effect on cost in all L - S combinations, particularly in $L2$ and $L3$, where the ranges of cost are over 100% larger than that of the more dispersed $L1$ (Fig. 2). However, ranges and IQRs of n_{req} are similar across L levels, suggesting that design choices have less potential impact on n_{req} in different landscape types.

Stand-scale tree pattern S also has a significant impact on inventory cost (Table 1). As with landscape pattern, average cost and n_{req} increase with increasing level of stand pattern aggregation ($S1$ - $S3$). Average cost generally increases linearly; however, there is a much larger (20%) increase in mean n_{req} between $S2$ and $S3$ compared to that between $S1$ and $S2$ (4%). Within each L level, ranges and IQRs of cost are similar across levels of S , but those of n_{req} are much larger for $S3$ (Fig. 2). This suggests that cluster plot design decisions have similar potential impact on inventory cost in different stand pattern types, but have very different potential impacts on n_{req} , particularly for the aggregated stand type. This is the opposite of what results suggest for landscape pattern aggregation, where cost variability is more sensitive than that of n_{req} to landscape pattern differences.

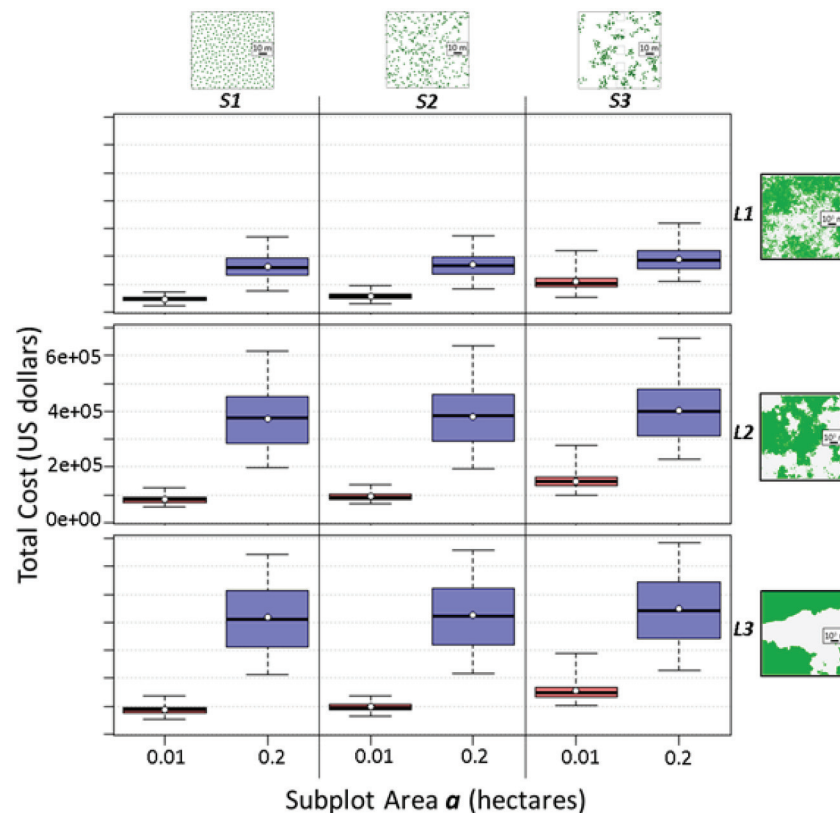
Impacts of plot design factors on costs and n_{req}

Of the three plot design variables evaluated, subplot area a was the most impactful, with the mean cost of 0.01 ha subplots being 71% lower than that of 0.2 ha subplots (Fig. 3). Compared to smaller subplots, ranges and IQRs of cost were much larger for 0.2 ha subplots, indicating that varying the other plot design factors can have high impacts when designing inventories with large subplots.

For the smaller subplots size (0.01 ha), increasing stand aggregation (from $S1$ to $S3$), lead to a high cost range and IQR. This increase in cost range and IQR in more aggregated stand patterns is of less magnitude when subplot size was larger (0.2 ha). This trend is reversed as landscape aggregation increases (from $L1$ to $L3$); cost variability was relatively stable for small subplots but increased dramatically between $L1$ and $L3$ for large subplots (Fig. 3). These findings are reinforced by the highly significant interaction between a and S and between a and L (Table 1) and highlight the complex relationships between design factors and attribute spatial patterns.

Subplot area a and number of subplots m interact significantly (Table 1), which was expected as they determine total plot area. The highest mean cost plot design was the one with the largest total plot area ($a = 0.2$, $m = 5$) while the lowest mean cost was for the nearly smallest total plot area ($a = 0.01$, $m = 3$) (Fig. 4). The situation is reversed for n_{req} ; the design with highest n_{req} was for the smallest plot ($a = 0.01$, $m = 5$), while the lowest n_{req} was for the largest plot ($a = 0.2$, $m = 5$) (Fig. 5). To summarize, results indicate that as a general pattern, increasing number of subplots and subplot area (and thus total plot area) increases cost and decreases n_{req} (Figs. 4 and 5).

Fig. 3. Boxplot representing distribution of total inventory cost values for each combination of subplot area (*a*), landscape type (*L1*, *L2*, *L3*), and stand type (*S1*, *S2*, *S3*). Shaded boxes are interquartile ranges, dark lines are medians, small circles are means, and dashed lines are range of values for each level. [Colour online.]



The effect of number of subplots on cost and n_{req} are also mediated by the heterogeneity type and level (Figs. 4 and 5); this is supported by the significant $a:m:L$ and $a:m:S$ interactions from the model results (Table 1). As landscape aggregation increases from *L1* to *L3*, the cost impact of raising the number of subplots increases, dramatically so for larger subplot sizes (Fig. 4). The same pattern of cost increase is observed as stand aggregation level increases, however, to a lesser extent. This suggests that cluster plot design decisions about number and size of subplots should thus be made by focusing more on the coarser-scale landscape pattern structure than on fine-scale stand patterns.

The above-described patterns of cost are reversed for n_{req} ; increasing the number of subplots has a larger impact on n_{req} as stand-scale aggregation increases from *S1* to *S3*, dramatically so for smaller subplot sizes (Fig. 5). The same pattern of higher impact of *m* is observed as landscape aggregation level increases, however, to a lesser extent. The opposing results when evaluating cost versus n_{req} highlight the importance of including both cost analyses and precision in design decisions.

Subplot separation (*d*) is the least impactful of the design variables studied. Although *d* occurs in several significant interactions in our model (Table 1), the practical effect of varying *d* is relatively small (Figs. 4 and 5). In general, as distance between subplots increases, the cost and the n_{req} decrease. However, the magnitude of this effect generally depends on the size of the subplots: for large subplots, increasing *d* has a larger impact on cost (Fig. 4) and a smaller impact on n_{req} (Fig. 5).

Discussion

This study uncovers the nature of the effects of spatial heterogeneity on inventory plot design efficiency, highlighting that

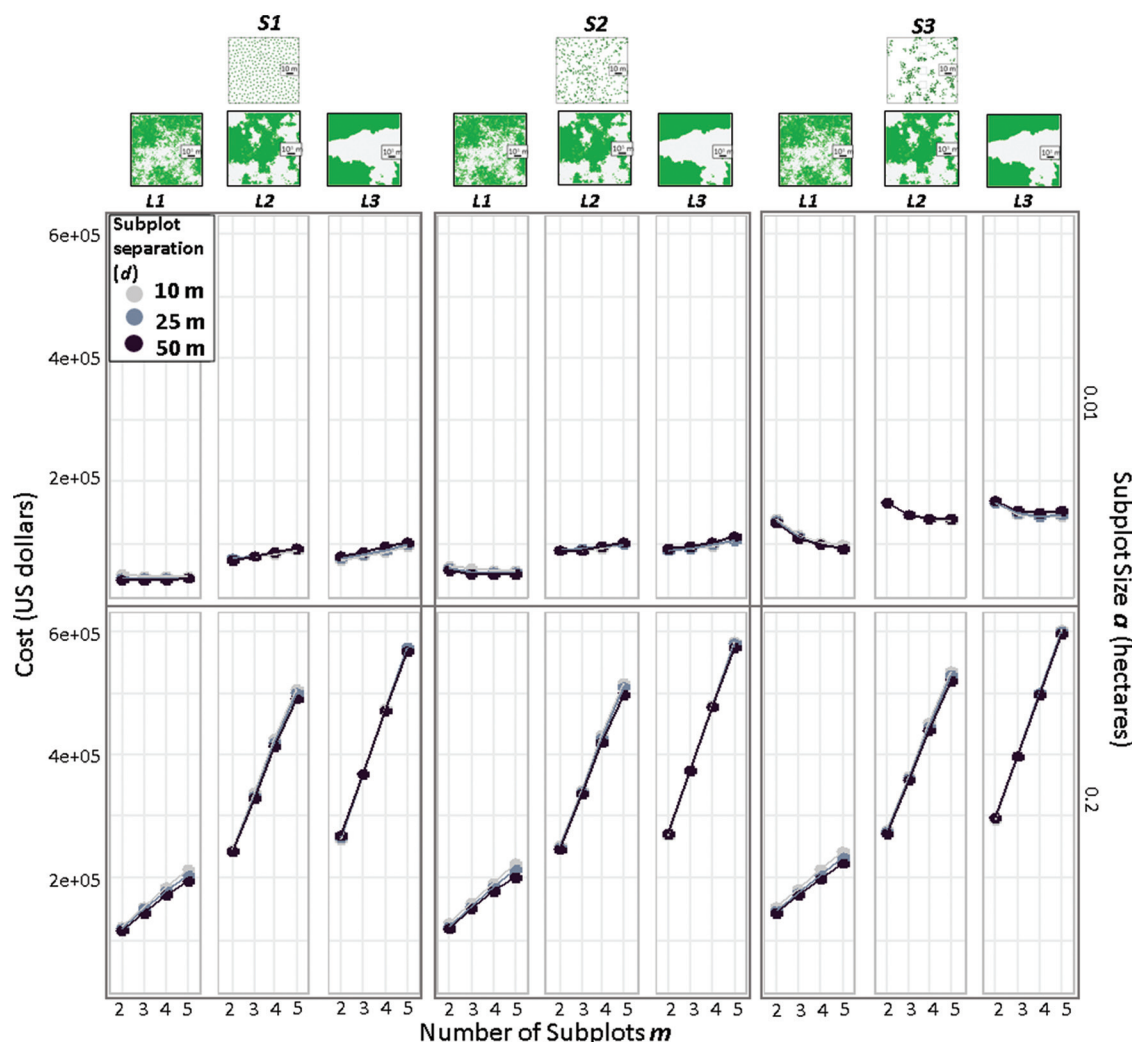
decisions related to plot design should be considered in the light of landscape- and stand-level heterogeneity. Depending on the type and level of heterogeneity, the decisions regarding subplot size, number, and separating distance have different importance.

Other studies have used either mapped or simulated stands to assess the effects of spatial pattern at the plot scale (Henttonen and Kangas 2015; Häbel et al. 2019). Our study is unique because we deliberately construct a broad range of spatial patterns at two scales, using tools offered by point process (Baddeley et al. 2015) and landscape pattern modelling (Gardner 1999), that reflect very different forest heterogeneity scenarios (Fig. 1). As in large-scale inventories that seek to monitor conversions to and from forest, like that of the USDA Forest Service's FIA program (USDA 2020), plots in our study often fell partially or entirely within nonforest gaps, contributing significantly to the variance of estimates, calling into question design guidelines constructed with results from studies in homogeneous areas.

Impacts of landscape heterogeneity *L* and stand heterogeneity *S* on costs

Figure 2 shows that there are important differences between mean costs across *L* levels, and smaller differences in mean cost across *S* levels. Both landscape- and stand-level pattern aggregation raise inventory costs. Less-aggregated landscapes (*L1*) are on average the least costly to inventory, and these costs are relatively insensitive to plot design choices. This is shown by how a broad range of n_{req} leads to a much narrower range of cost relative to the more aggregated landscape pattern types. For intermediate *L2* and aggregated *L3* landscapes, this situation is different; there are generally broad ranges of cost differences relative to the range of required sample sizes, indicating that design choices

Fig. 4. Interaction plots showing the relationship between mean total inventory cost and different levels of the factors used in the simulation (*a*, subplot area; *d*, subplot edge separation distance; *m*, number of subplots; *L*, landscape type; *S*, stand type). [Colour online.]



affecting CV and thus n_{req} are more important in landscapes with aggregated forest patterns.

Uniform stands (*S1*) are on average the least costly to inventory, followed by stands with random tree patterns (*S2*); both are less costly than inventories in stands with clustered patterns (*S3*) (Fig. 2). Uniform and random tree patterns have similar cost– n_{req} relationships, as we expected, because as analysis windows get larger, point density estimates from random patterns tends to converge with those obtained from uniform patterns much faster than do those from aggregated patterns (Wiegand and Moloney 2014). At the scale of plots (i.e., analysis windows) and with the average tree densities used in this experiment, means from the random pattern resemble those of the uniform pattern. Had we used lower tree densities, or smaller plots, differences would have been greater.

Our results align with those of other studies. In a study using a sampling simulation on a set of stem-mapped tropical forest stands, Picard et al. (2018) found that spatial pattern variations due to differing forest types and stand structures had strong impacts on the CV – plot size relationships and cost optimality of different designs, with coarse-scale heterogeneity effects being stronger than those of local-scale heterogeneity. Similarly, Baraloto et al. (2013) found meaningful differences in plot design – variance relationships across five neotropical forests, and much smaller differences within each forest. Häbel et al. (2019) found that costs

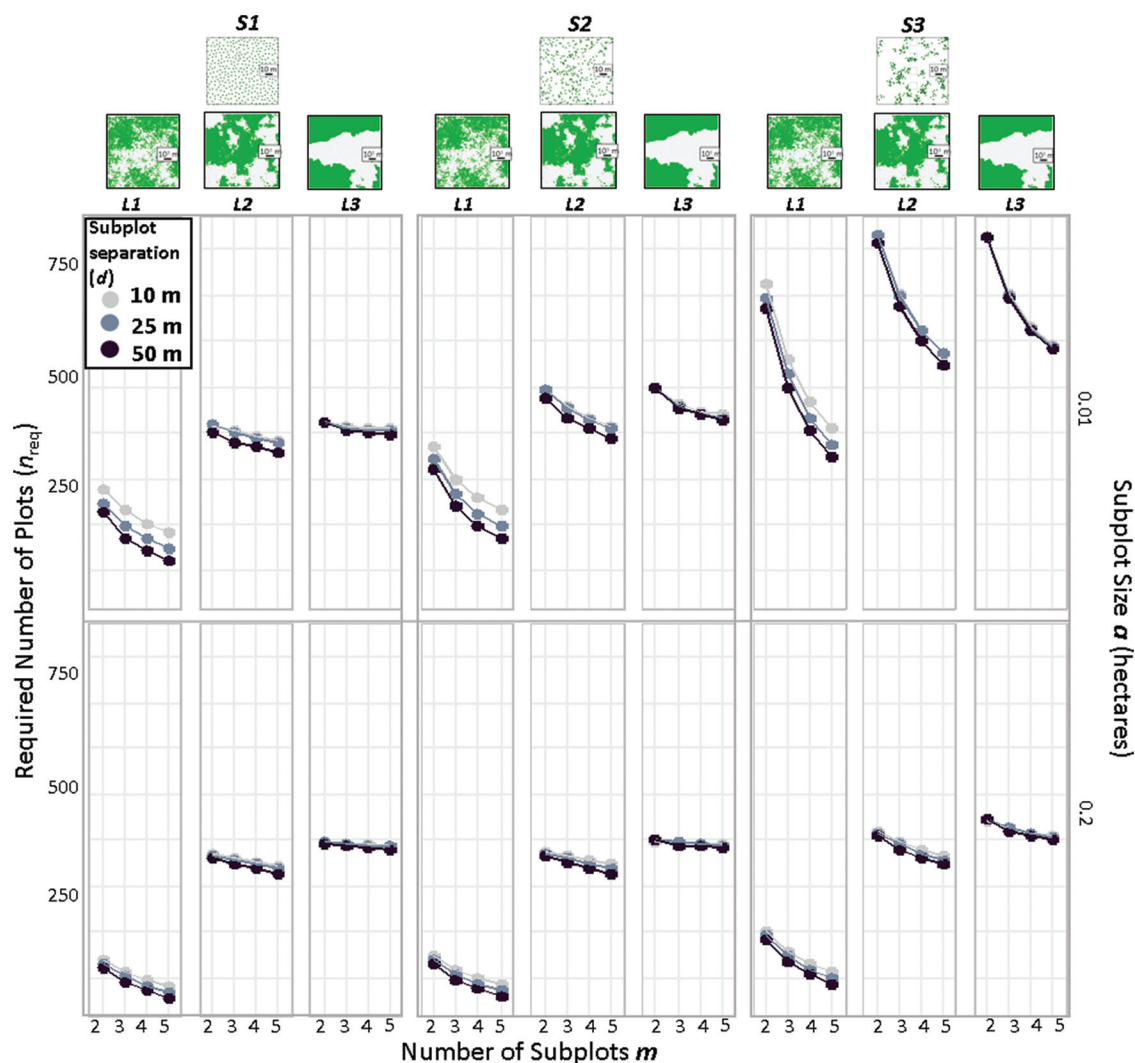
are lowest in more regular stands than in clustered or random stands, and that increasing subplot size had a larger impact on precision in clustered patterns than in homogeneous patterns. Clearly, these findings depend upon the nature of the heterogeneity introduced at each scale; sharp forest – nonforest boundaries will have a greater variance-inducing effect than will finer-scale variations in levels of forest attributes that do not create gaps into which entire plots or subplots fall.

Impacts of interactions among design factors, *L*, and *S* on costs

Effects of *m* and *a*

The overall pattern of change in cost and n_{req} with plot area is largely due to the well-known inverse exponential relationship between plot area and CV (and thus n_{req}) identified by Smith (1938) and several others; as plots get larger, there is a diminishing rate of reduction in variance. However, increasing plot size incurs costs, and the results indicate that, depending upon *L*–*S* type, the CV and n_{req} reduction benefits of increasing plot size do not balance out the additional field costs incurred when plots get very large. This finding motivates a recommendation to use clusters that are significantly smaller than 1 ha, which are commonly recommended, when cost function parameters and attribute variance structure resemble those from our experiment.

Fig. 5. Interaction plots showing the relationship between mean required number of plots (n_{req}) to meet allowable error (AE) and different levels of the factors used in the simulation (a = subplot area, d = subplot edge separation distance, m = number of subplots, L = landscape type, S = stand type). [Colour online.]



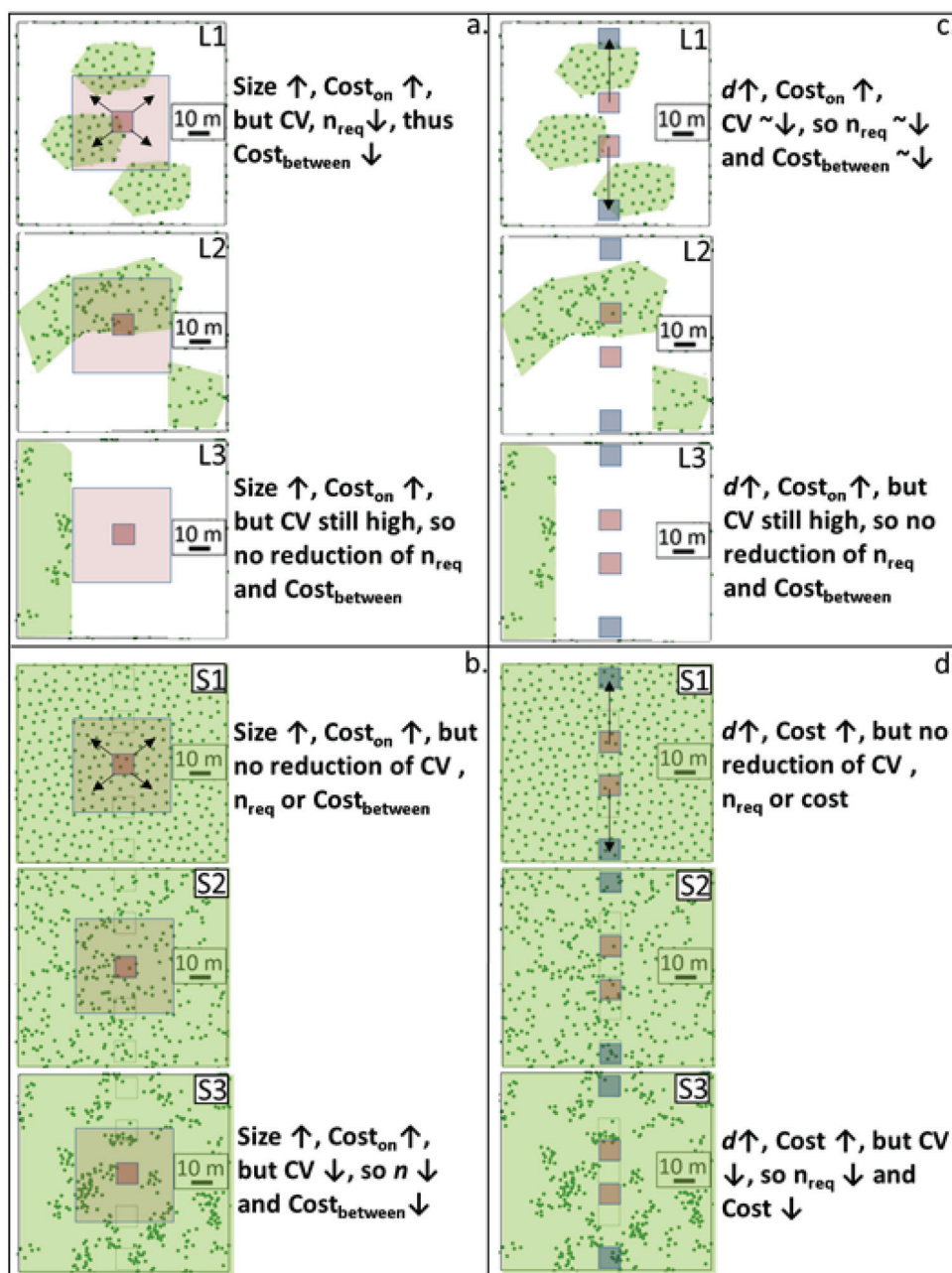
Several authors have found that cost optimization favoured smaller plots ranging in size from 0.07 to 0.2 ha (Mesavage and Grosenbaugh 1956; Scott 1993; Yim et al. 2015; Lynch 2017b). Others have found that in areas with extreme levels of between-plot costs, or to meet needs like measuring site-based diversity or calculating stand indices that require more trees to be meaningful, plots between 0.25 and 2 ha might be required (Zenner and Peck 2009; Rejou-Mechain et al. 2014; Picard et al. 2018). In addition, if inventory plots are to be linked to remote sensing data and used to train models, larger plots help mitigate the effects of spatial mismatch between plots and pixels (Næsset et al. 2015; GFOI 2016).

Given this broad range of recommendations, it becomes clear that studies like this one, that not only assess different designs in a methodical way, but also experimentally integrate different heterogeneity scenarios, can help reveal mechanisms behind cost efficiency. For example, the results of our study indicate that as landscape aggregation level and thus forest patch sizes increase, cost – plot size relationships vary in alignment with principles identified in the conceptual model shown in Fig. 6a. In the more-aggregated $L3$ landscape type, increasing subplot size is less likely to cause the plot to represent a microcosm of

the larger study area than it is in $L1$, so while costs on the plot increase, n_{req} and cost between plots does not decrease as much as it does in $L1$ (Figs. 4, 5, 6a). If each plot's status is more similar to that of the landscape as a whole, as in $L1$, its contribution to variance and CV will be smaller, leading to, on average, fewer and (or) less costly plots needed to achieve AE requirements from the sample.

For stand-scale aggregation S , similar principles apply. At higher stand aggregation levels ($S3$), Fig. 4 reveals an important phenomenon: a negative exponential pattern of the relationship between cost and number of subplots in the panel associated with smaller subplots, and a positive relationship in the panel associated with larger subplots. This suggests that the optimal (lowest cost) plot size lies somewhere between 0.05 ha (five 0.01 ha subplots) and 0.4 ha (two 0.2 ha subplots). This supports the hypothesis that increasing plot size in aggregated stands makes each plot more like a microcosm of the stand-scale patterns found on the landscape than it does in the stands where trees are more evenly distributed like $S1$ and $S2$ (Fig. 6b), where the optimum plot size generally appears to be at or below 0.01 ha. For both L and S , there is therefore little cost advantage to increasing plot size unless the heterogeneity type is such that

Fig. 6. Conceptual model showing hypothesized mechanisms for how increasing plot size and separation affect coefficient of variation (CV) and costs across levels of landscape type (*L*) (a, c) and stand type (*S*) (b, d). Arrows represent increasing of plot size or separation distance from smaller (subplot area (*a*) = 0.01 ha) or closer (subplot separation distance (*d*) = 10 m) to larger (*a* = 0.2 ha) or more separated (*d* = 50 m) cluster plot designs. Green areas represent forested areas; small dots represent tree locations. [Colour online.]



reductions in CV and subsequent savings in *n*_{req} offset the on-plot cost increases from measuring larger plots.

Effects of *d*

Separating subplots has a similar effect on variance as does increasing plot area, but for slightly different reasons. In theory, separating subplots should lead to the mean landscape condition being captured more efficiently for the same total plot area because the cost inefficiencies associated with measuring redundant information in spatially autocorrelated patches should be mitigated (Cochran 1977; Kleinn 1994; Tokola and Shrestha 1999; Yim et al. 2015). In practice, the magnitude of the impacts of

subplot separation distance on CV and cost is linked to the relationship between it and the spatial structure of the resource under study and the relative costs associated with walking between subplots.

Increasing landscape aggregation (from *L1* to *L3*) leads to a reduction of cost benefits from subplot separation (Fig. 4). As with increasing plot size, separating subplots is more likely to capture the variability found on the landscape, reducing variance, when forest patches are more dispersed, like in *L1*. However, separating subplots found in large, aggregated forest or nonforest patches (*L2* and *L3*) will lead to less reduction of variance and in *n*_{req}, but higher on-plot costs from additional walking. This will tip the balance and favour lower values of *d* in more aggregated landscapes (Fig. 6c). As levels of

stand-scale heterogeneity increase (*S1* to *S3*), subplot separation becomes more cost-advantageous, because in more aggregated stands, subplots move into or out of patches of trees over relatively short distances (Fig. 6d).

Additional considerations

We used approaches similar to others who have built cost models for forest inventories (Zeide 1980; Scott 1993; Zhang et al. 1994; Yim et al. 2015; Yang et al. 2017), and attempted to provide estimates of cost parameters that were reasonable based on personal experience and discussions with field inventory personnel. However, there are several factors for which we did not account. First and foremost, our cost model is based on the use of fixed area cluster plot designs with the cost parameters we selected. These model parameters would vary if different field techniques (lasers for distance and height measurements versus tapes and clinometers, or variable radius versus fixed area plots) were used. Yang et al. (2017) and Lynch (2017b) report that for the datasets they examined, there was a rather flat optimality region associated with the cost (or relative standard error): plot size curve, suggesting that excessive fine tuning of cost and design parameters might not be important in practice. However, Chen et al. (2019) found a much narrower range of optimality, suggesting that under certain heterogeneity scenarios and when costs are a limiting factor, design decisions become more important. Due to the varied influences of inventory methods and attribute heterogeneity on costs, we designed our modelling framework and simulation experiment with enough flexibility to allow for a wide variety of spatial heterogeneity scenarios and user-specified cost model parameters.

Another factor that could alter the results of our study is the choice of the attribute on which to optimize. Yang et al. (2017) employ a variable radius plot procedure that focuses sampling effort on tree counts because, particularly for smaller diameter classes, this attribute is generally more variable than volume or basal area. We chose *N* because, for trees above 12.7 cm in diameter, its variance is intermediate to that of *G* or *V* (USDA 2020) and because point pattern simulation of *N* is much less complex than alternative approaches that include effects of tree size on spatial patterns. In real-world inventories, there are multiple objectives and types of information needs, and it is important to note that different attribute choices will have different variance structures, affecting outcomes of plot design choices.

Finally, the logistics of conducting the inventory are crucial determinants of cost efficiency. In the US Forest Service's national forest inventory, for example, plots are pre-screened with aerial imagery before a field crew is sent. In the context of our experiment, this would dramatically reduce the impact of n_{req} , as many plots that fall entirely within nonforest patches would not require a field visit. The travel logistics are also key; for example, with respect to plot size, higher between-plot travel costs would tend to shift the optimum toward larger plots, as on-plot costs would represent a lower proportion of overall inventory cost (Zeide 1980).

Our research provides practical guidance to decision-makers. Firstly, the results strongly suggest that a simulation study aimed at modelling the effects of spatial pattern type and scale, plot design factors, and AE requirements on inventory cost can be extremely valuable, as our results show that there are important total inventory cost differences across the range of plot design choices that are often considered. Traditional paradigms of inventory plot design tend to rely on large inventory plots, with the idea that these somehow "capture" rare events or give more accurate estimates of the status of the attribute. What is not considered, however, is that a common goal of an efficient inventory design is not to maximize the unknowable quantity "accuracy" but rather to meet AE requirements for minimum cost. We have therefore provided a conceptual framework, as well as code for conducting simulations, with which practitioners can approach inventory plot design choices in ways that are better than simply resorting to the "bigger is better" orthodoxy.

Secondly, we have demonstrated the value of using the tools of point process modelling and landscape simulation for inventory plot design investigations. The importance of the effects of plot design on inventory costs changes with spatial heterogeneity scenario, and while our specific heterogeneity scenarios are not perfect facsimiles of nature, the principles that our simulation experiment reveal as well as our methods and software code can be used to guide plot design experiments in the field or in future simulation studies. For example, to simulate tree point patterns that resemble those found in a particular study area, a practitioner can either assume that the dominant pattern type can be described by one of the three stand-scale pattern scenarios used in this study (Fig. 1) or implement a more sophisticated point process modelling approach using data from stem-mapped plots or stands (Lister and Leites 2018). For example, Henttonen and Kangas (2015) and Häbel et al. (2019) found that plot design optimality varied greatly by combination of inventory attribute and tree size distribution-induced spatial pattern properties.

With respect to landscape patterns, our recommendation for practitioners would be to test plot designs using existing forest–nonforest maps in place of simulated maps, because these exist at the appropriate resolution for most areas. For example, by integrating realistic, simulated point pattern maps with realistic land cover maps created from actual landscapes, a reasonable facsimile of the population could be produced and various plot types overlaid and tested using the analysis framework we have presented.

Conclusions

The type and scale of heterogeneity found in the population being inventoried matters when making plot design choices. In our experiment, landscape-scale heterogeneity affects the magnitude of costs more than stand-scale heterogeneity, with costs increasing with aggregation level. At the stand scale, areas with clustered tree patterns are more costly to survey than those with regular tree patterns, with random tree patterns being intermediate to the two. It is more advantageous to increase plot size in stands with aggregated tree patterns compared to stands with more uniform tree patterns, mainly because there is a larger reduction in CV as plot size, and to a lesser extent subplot separation distance, go up. These relationships are sensitive to the scale of heterogeneity of the attribute of interest, relative to the dimensions of the plot. Subplot count and area are more important design decisions affecting total cost than subplot separation distance, although once plot area has been chosen, separation distance matters, with smaller distances being more cost-effective for smaller plots, and larger distances being more cost-effective for larger plots. This relationship is affected more by landscape than by stand heterogeneity; there is less cost penalty from subplot separation in more uniform landscapes.

This study is unique in that it recognizes inventory situations where the paradigm includes a survey of the entire land base including nonforest areas. We designed a simulation experiment framework that allows testing different plot design factor combinations using modelling and simulation tools from the disparate fields of point process modelling and landscape ecology. This study not only reveals how various plot design factors interact under different heterogeneity scenarios to affect costs, but it also provides a conceptual framework and computer code with which to examine these interactions. Our goal was to provide guidance to practitioners and tools that can serve as a foundation for future studies in this important area.

Acknowledgements

We thank Ephraim Hanks, Eric Zenner, Doug Miller, Charles Scott, James Westfall, Samidha Shetty, Tyler Garner, John Stanovick, and two anonymous reviewers for advice and guidance on earlier versions of this manuscript. This work was partially supported by the USDA

National Institute of Food and Agriculture and Hatch Appropriations under project No. PEN04700 and accession No. 1019151.

References

- Arvanitis, L.G., and O'Regan, W.G. 1967. Computer simulation and economic efficiency in forest sampling. *Hilgardia*, **38**(2): 133–164. doi:10.3733/hilg.v38n02p133.
- Baddeley, A., and Turner, R. 2005. spatstat: an R package for analyzing spatial point patterns. *J. Stat. Softw.* **12**(6): 1–42. doi:10.18637/jss.v012.i06.
- Baddeley, A., Rubak, E., and Turner, R. 2015. Spatial point patterns — methodology and applications with R. Chapman & Hall/CRC Interdisciplinary Statistics, New York.
- Baraloto, C., Molto, Q., Rabaud, S., Herault, B., Valencia, R., Blanc, L., et al. 2013. Rapid simultaneous estimation of aboveground biomass and tree diversity across neotropical forests: a comparison of field inventory methods. *Biotropica*, **45**(3): 288–298. doi:10.1111/btp.12006.
- Bates, D., Mächler, M., Bolker, B., and Walker, S. 2015. Fitting linear mixed-effects models using lme4. *J. Stat. Softw.* **67**(1): 1–48. doi:10.18637/jss.v067.i01.
- Bechtold, W.A., and Patterson, P.L. 2005. The enhanced forest inventory and analysis program — national sampling design and estimation procedures, GTR-SRS-80. Gen. Tech. Rep., US Department of Agriculture, Forest Service, Southern Research Station, Asheville, NC. Available from <https://www.fs.usda.gov/treesearch/pubs/20371> [accessed 27 November 2020].
- Brink, G.E., and Schreuder, H.T. 1992. ONEPHASE: a simulation program to compare 1-phase sampling strategies. Research Paper, USDA Forest Service, Rocky Mountain Forest and Range Experiment Station, Fort Collins, Colorado. Available from <http://catalog.hathitrust.org/Record/011397236>.
- Chen, Y., Yang, T.-R., Hsu, Y.-H., Kershaw, J.A., and Prest, D. 2019. Application of big BAF sampling for estimating carbon on small woodlots. *For. Ecosyst.* **6**(1): 13. doi:10.1186/s40663-019-0172-4.
- Cochran, W. 1977. Sampling techniques. John Wiley & Sons, New York.
- FAO. 2017. Voluntary guidelines on national forest monitoring. Food and Agriculture Organization of the United Nations (FAO), Rome. Available from <http://www.fao.org/3/a-i6767e.pdf>.
- Freese, F. 1962. Elementary forest sampling. USDA Forest Service, Washington, D.C. Available from <https://books.google.com/books?id=wilkwAAAAAYAJ>.
- Gardner, R.H. 1999. RULE: Map generation and a spatial analysis program. In *Landscape ecological analysis: Issues and applications*. Edited by J.M. Klopatek and R.H. Gardner. Springer, New York, NY. pp. 280–303. doi:10.1007/978-1-4612-0529-6_13.
- Gardner, R.H. 2017. Characterizing categorical map patterns using neutral landscape models. In *Learning landscape ecology: a practical guide to concepts and techniques*. Edited by S.E. Gergel and M.G. Turner. Springer, New York, NY. pp. 83–103. doi:10.1007/978-1-4939-6374-4_6.
- GFOI. 2016. Integration of remote-sensing and ground-based observations for estimation of emissions and removals of greenhouse gases in forests: Methods and guidance from the Global Forest Observations Initiative. Food and Agriculture Organization, Rome. Available from <https://www.fs.usda.gov/treesearch/pubs/56461> [accessed 27 November 2020].
- Gove, J. 2017. Some refinements on the comparison of areal sampling methods via simulation. *Forests*, **8**(10): 393. doi:10.3390/f8100393.
- Häbel, H., Kuronen, M., Henttonen, H.M., Kangas, A., and Myllymäki, M. 2019. The effect of spatial structure of forests on the precision and costs of plot-level forest resource estimation. *For. Ecosyst.* **6**(1): 8. doi:10.1186/s40663-019-0167-1.
- Hasler, M., and Hornik, K. 2007. TSP — Infrastructure for the traveling salesperson problem. *J. Stat. Softw.* **23**(2): 1–21. doi:10.18637/jss.v023.i02.
- Henttonen, H., and Kangas, A. 2015. Optimal plot design in a multipurpose forest inventory. *For. Ecosyst.* **2**: 31. doi:10.1186/s40663-015-0055-2.
- Hijmans, R.J. 2019. Raster: Geographic data analysis and modeling, version 2.9-23. Available from <https://CRAN.R-project.org/package=raster> [accessed 27 November 2020].
- Hou, Z., Xu, Q., Hartikainen, S., Antilla, P., Packalen, T., Maltamo, M., and Tokola, T. 2015. Impact of plot size and spatial pattern of forest attributes on sampling efficacy. *For. Sci.* **61**(5): 847–860. doi:10.5849/forsci.14-197.
- IPCC. 2006. IPCC guidelines for national greenhouse gas inventories. Institute for Global Environmental Strategies, Japan. Available from <http://www.ipcc-nggip.iges.or.jp/public/2006gl/> [accessed 27 November 2020].
- Kershaw, J.A., Richards, E.W., McCarter, J.B., and Oborn, S. 2010. Spatially correlated forest stand structures: a simulation approach using copulas. *Comput. Electron. Agr.* **74**(1): 120–128. doi:10.1016/j.compag.2010.07.005.
- Kleinn, C. 1994. Comparison of the performance of line sampling to other forms of cluster sampling. *For. Ecol. Manag.* **68**(2–3): 365–373. doi:10.1016/0378-1127(94)90057-4.
- Kuznetsova, A., Brockhoff, P.B., and Christensen, R.H.B. 2017. lmerTest package: tests in linear mixed effects models. *J. Stat. Softw.* **82**(13): 1–26. doi:10.18637/jss.v082.i13.
- Larson, A., Lutz, J., Donato, D., Freund, J., Swanson, M., HilleRisLambers, J., et al. 2015. Spatial aspects of tree mortality strongly differ between young and old-growth forests. *Ecology*, **96**(11): 2855–2861. doi:10.1890/15-0628.1.
- Lister, A.J., and Leites, L.P. 2018. Modeling and simulation of tree spatial patterns in an oak-hickory forest with a modular, hierarchical spatial point process framework. *Ecol. Model.* **378**: 37–45. doi:10.1016/j.ecolmodel.2018.03.012.
- Loetsch, F., and Haller, K.E. 1973. Forest inventory. BLV, Munich, Germany.
- Lynch, A.M. 2003. Comparison of fixed-area plot designs for estimating stand characteristics and western spruce budworm damage in southwestern U.S.A. forests. *Can. J. For. Res.* **33**(7): 1245–1255. doi:10.1139/x03-044.
- Lynch, T.B. 2017a. Optimal plot size or point sample factor for a fixed total cost using the Fairfield Smith relation of plot size to variance. *For. Int. J. For. Res.* **90**(2): 211–218. doi:10.1093/forestry/cpw038.
- Lynch, T.B. 2017b. Optimal sample size and plot size or point sampling factor based on cost-plus-loss using the Fairfield Smith relationship for plot size. *Forestry*, **90**(5): 697–709. doi:10.1093/forestry/cpx024.
- Mackisack, M.S., and Wood, G.B. 1990. Simulating the forest and the point-sampling process as an aid in designing forest inventories. *For. Ecol. Manag.* **38**(1): 79–103. doi:10.1016/0378-1127(90)90087-R.
- Mesavage, C., and Grosenbaugh, L.R. 1956. Efficiency of several cruising designs on small tracts in north Arkansas. *J. For.* **54**(9): 569–576. doi:10.1093/jof/54.9.569.
- Næsset, E., Bollandsås, O.M., Gobakken, T., Solberg, S., and McRoberts, R.E. 2015. The effects of field plot size on model-assisted estimation of aboveground biomass change using multitemporal interferometric SAR and airborne laser scanning data. *Remote Sens. Environ.* **168**: 252–264. doi:10.1016/j.rse.2015.07.002.
- Oehlert, G.W. 2000. A first course in design and analysis of experiments. W.H. Freeman, New York.
- Pebesma, E.J., and Bivand, R.S. 2005. Classes and methods for spatial data in R. *R News*, **5**(2): 9–13.
- Picard, N., Nouvellet, Y., and Sylla, M.L. 2004. Relationship between plot size and the variance of the density estimator in West African savannas. *Can. J. For. Res.* **34**(10): 2018–2026. doi:10.1139/x04-079.
- Picard, N., Gamarra, J.G.P., Birigazzi, L., and Branthomme, A. 2018. Plot-level variability in biomass for tropical forest inventory designs. *For. Ecol. Manag.* **430**: 10–20. doi:10.1016/j.foreco.2018.07.052.
- R Core Team. 2018. R: A Language and Environment for Statistical Computing, version 3.5.1. R Foundation for Statistical Computing, Vienna, Austria. Available from <https://www.R-project.org/> [accessed 27 November 2020].
- Reich, R.M., and Arvanitis, L.G. 1992. Sampling unit, spatial distribution of trees, and precision. *North. J. Appl. For.* **9**(1): 3–6. doi:10.1093/njaf/9.1.3.
- Rejou-Mechain, M., Muller-Landau, H., Detto, M., Thomas, S., Le Toan, T., Saatchi, S., et al. 2014. Local spatial structure of forest biomass and its consequences for remote sensing of carbon stocks. *Biogeosciences*, **11**(23): 6827–6840. doi:10.5194/bg-11-6827-2014.
- Riitters, K., Wickham, J., Costanza, J., and Vogt, P. 2016. A global evaluation of forest interior area dynamics using tree cover data from 2000 to 2012. *Landscape Ecol.* **31**(1): 137–148. doi:10.1007/s10980-015-0270-9.
- Rosenkrantz, D.J., Stearns, R.E., and Lewis, P.M. 2009. An analysis of several heuristics for the traveling salesman problem. In *Fundamental problems in computing: Essays in honor of Professor Daniel J. Rosenkrantz*. Edited by S.S. Ravi and S.K. Shukla. Springer Netherlands, Dordrecht. pp. 45–69. doi:10.1007/978-1-4020-9688-4_3.
- Schreuder, H.T., Banyard, S.G., and Brink, G.E. 1987. Comparison of three sampling methods in estimating stand parameters for a tropical forest. *For. Ecol. Manag.* **21**(1): 119–127. doi:10.1016/0378-1127(87)90076-4.
- Scott, C.T. 1993. Optimal design of a plot cluster for monitoring. In *Proceedings of the IUFORO S.4.11 Conference*. Edited by K. Rennolls and G. Gertner. University of Greenwich, London, UK. pp. 233–242. Available from <http://www.nrs.fs.fed.us/pubs/42932> [accessed 27 November 2020].
- Smith, H. 1938. An empirical law describing heterogeneity in the yields of agricultural crops. *J. Agric. Sci.* **28**: 1–23. doi:10.1017/S0021859600050516.
- Stoyan, D., and Penttinen, A. 2000. Recent applications of point process methods in forestry statistics. *Stat. Sci.* **15**(1): 61–78. doi:10.1214/ss/1009212674.
- Thompson, S.K. 2012. Sampling. 3rd ed. John Wiley and Sons, Inc., Hoboken, NJ.
- Tokola, T., and Shrestha, S.M. 1999. Comparison of cluster-sampling techniques for forest inventory in southern Nepal. *For. Ecol. Manag.* **116**(1–3): 219–231. doi:10.1016/S0378-1127(98)00457-5.
- USDA. 2020. Evaluator retrieval for volume per acre of forest, basal area per acre of forest, and trees per acre of forest for trees greater than 5 inches dbh, Pennsylvania EVALID 422018. Forest Inventory and Analysis, U.S. Department of Agriculture, Forest Service, Northern Research Station. Available from <http://apps.fs.usda.gov/Evaluator/evaluator.jsp> [accessed 27 November 2020].
- Wiant, H.V., and Yandle, D.O. 1980. Optimum plot size for cruising sawtimber in eastern forests. *J. For.* **78**(10): 642–643. doi:10.1093/jof/78.10.642.
- Wiegand, T., and Moloney, K.A. 2014. Handbook of spatial point-pattern analysis in ecology. CRC Press, Taylor & Francis Group, Boca Raton, FL.
- Yang, T.-R., Hsu, Y.-H., Kershaw, J.A., McGarrigle, E., and Kilham, D. 2017. Big BAF sampling in mixed species forest structures of northeastern North America: Influence of count and measure BAF under cost constraints. *Forestry*, **90**(5): 649–660. doi:10.1093/forestry/cpx020.
- Yim, J.-S., Shin, M.-Y., Son, Y., and Kleinn, C. 2015. Cluster plot optimization for a large area forest resource inventory in Korea. *For. Sci. Technol.* **11**(3): 139–146. doi:10.1080/21580103.2014.968222.
- Zarnoch, S.J., and Bechtold, W. 2000. Estimating mapped-plot forest attributes with ratios of means. *Can. J. For. Res.* **30**(5): 688–697. doi:10.1139/cjfr-30-5-688.
- Zeide, B. 1980. Plot size optimization. *For. Sci.* **26**(2): 251–257. doi:10.1093/forestscience/26.2.251.
- Zenner, E.K., and Peck, J.E. 2009. Characterizing structural conditions in mature managed red pine: spatial dependency of metrics and adequacy of plot size. *For. Ecol. Manag.* **257**(1): 311–320. doi:10.1016/j.foreco.2008.09.006.

Zhang, R., Warrick, A.W., and Myers, D.E. 1994. Heterogeneity, plot shape effect and optimal plot size. *Geoderma*, 62(1-3): 183-197. doi:10.1016/0016-7061(94)90035-3.

Appendix A. Cost model and explanation of cost components used to calculate total inventory cost from specific plot design-required sample size scenarios

Cost model

The cost model used to calculate total inventory cost based on assumptions about crew salaries, work rates, travel rates, and daily logistical costs is described below. The combination of simulated forest pattern and cluster plot configuration leads to on-plot costs, and required sample sizes, along with assumptions about travel speed rates, affect between-plot costs. Daily rates are affected by time required to complete the entire inventory. Assumptions of the cost model are as follows:

1. Plots are visited in succession, with field crews receiving salary for 8 h/day. At the end of each day, the field crew will lodge at a negligible distance from the route between plots.
2. Logistical costs, such as training, equipment purchases, lodging costs, vehicle purchase or lease, and overhead costs, can be subsumed into a daily cost, with the assumption that they can be amortized using the time period at which the inventory field data are being collected.
3. Any other costs associated with on- or between-plot work can be subsumed into these components.
4. There is no upper time or cost limit to complete the inventory, as the goal is to seek the design that achieves AE for a single attribute of interest for the minimum cost.
5. A single crew is used to visit all plots, i.e., there are no crews working in parallel.
6. Several additional cost components are not considered in order to avoid excessive complexity, such as additional time required to measure edge trees (Zeide 1980), rest or data entry times, and so on.

The cost model with three terms, reflecting on-plot costs, between-plot costs, and logistical costs, is as follows:

$$(A1) \quad C_T = S_c \times n \times T_p + S_c \times [n \times T_t + (n - 1) \times T_d] + C_L \times D$$

where

C_T = total cost (US dollars),

S_c = total field crew salary (US dollars/h); this was set to \$100/h,

n = required number of plots to meet precision requirements, calculated from the CV and eq. 1,

T_p = time (h) required on-plot to complete plot (eq. A3),

T_d = time (h) required driving from one plot to the next (eq. A6),

T_t = transition time spent between leaving vehicle and arriving at plot, then, when finished with measurements, leaving the plot and entering vehicle (round trip time was set to 0.33 h/plot). Note that separating this factor, which could in principle be combined with driving time into a "between-plot cost" component, allows for experimental manipulation of the on-plot : between-plot cost ratio without having to adjust driving velocity.

Logistical cost components:

C_L = logistical costs (US dollars/day); this was set to \$200/day

D = number of days to complete inventory (eq. A2)

$$(A2) \quad D = [n \times (T_p + T_t) + (n - 1) \times T_d] / 8 \text{ h/day}$$

Note that logistical components, as we have incorporated them, are in effect daily costs that go up as more days are required to complete the inventory, given the number of required plots and time expenditure required for transportation and field work.

On-plot time components:

$$(A3) \quad T_p = (d_t + m \times d_c + d_m) / V_w + T_m \times p$$

where

d_t = distance walked along shortest path between trees on the plot, returning to first tree (m) (Fig. A1a). This shortest path distance was calculated using an algorithmic solution to what is known as the Traveling Salesperson Problem (TSP), farthest insertion method (Rosenkrantz et al. 2009) found within the R package TSP (Hasler and Hornik 2007). This approach was used to find a common method for tracing an efficient route between all of the trees found on a plot, regardless of the spatial pattern, and regardless of the fact that, in practice, field crews will seldom visit trees in this manner. d_t is thus considered to be a reasonable proxy for inter-tree walking time given that spatial patterns of trees differ greatly by level of S.

m = number of subplots.

d_c = distance from center of subplot to each of four corners and back to center (eq. A4); this is meant to be a simplistic representation of costs associated with marking the corners of a square subplot.

p = no. of trees on a plot

$$(A4) \quad d_c = 2 \times \{4 \times [0.5 \times \sqrt{2 \times l_s^2}]\}$$

where

l_s = length of one side of a subplot (Fig. A1b)

d_m = distance between subplot centers and back to subplot 1 (eq. A5)

$$(A5) \quad d_m = 2 \times \{[2 \times (0.5 \times l_s) + d] \times (m - 1)\}$$

where

d = separation distance between subplot sides (Fig. A1c)

V_w = walking velocity (m/h); this was set to 1609.5 m/h

T_m = time spent measuring a tree (h/tree); this was set to 0.05 h/tree; we ignore the fact that edge trees will require more time to discern and measure, or that many national forest inventories will require more measurement time per tree, particularly in dense stands if heights are measured.

Between-plot driving time components:

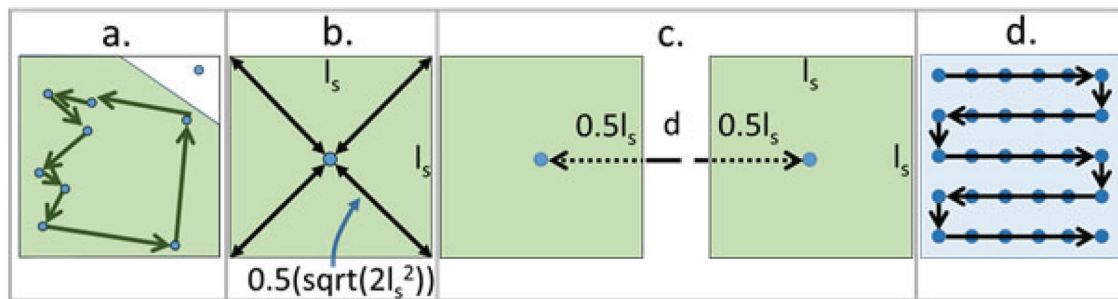
$$(A6) \quad T_d = [\sqrt{A}/\sqrt{n}] / V_d$$

where

A = area (km²) of the population; this was set to 10 000 km² for logistical reasons related to the computational requirements of the simulation, and to provide a large enough area over which to distribute a spatially balanced sample of 49 points that are assumed to be independent of one another.

V_d = velocity (km/h) of inter-plot travel; this was set to 48 km/h.

Fig. A1. Graphical depiction of components of the cost model used in the study. (a) Example of the Traveling Salesman Problem derived shortest distance between a set of forest trees (points) found in a forest (green, shaded areas in all graphics; nonforest trees are not measured) (d_t). (b) Example of the calculation of subplot establishment walking distance (d_c). (c) Example of the walking distance back and forth between two subplot centers (d_m). (d) A depiction of the assumed travel route and length of the set of segments that need to be traversed to visit all plots. l_s , length of one side of a subplot. [Colour online.]



Copyright of Canadian Journal of Forest Research is the property of Canadian Science Publishing and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.